

Expressibility Metrics

Variational Quantum Algorithms · Module 8 — Expressivity Vs Trainability

Node 8.2

Position in Knowledge Graph

System: Variational Quantum Algorithms **Structural Domain:** Expressivity vs Trainability **Node:** 8.2

Node 8.1 introduced expressivity as a *family* of related but inequivalent quantities. This node makes those quantities concrete: it walks through the **specific metrics** in the literature, how each one is computed, what each one tells you, and when each is the right tool.

The pragmatic motivation: if you are designing or comparing ansätze, you need a number — preferably one you can actually compute on a laptop. The pure theoretical view (reachable manifold dimension, Lie algebra dimension) doesn’t scale. The metrics in this node **do** scale, at the cost of capturing only a slice of the full picture.

What A Good Expressibility Metric Should Do

Before walking through the metrics, the desiderata:

- 1. Computability.** You should be able to compute it from samples (parameter draws + cost evaluations or fidelity calculations). No exact integrals, no exponential scaling.
- 2. Comparability.** Two ansätze with the same number of parameters and qubits should give comparable numbers. The metric shouldn’t be dominated by ansatz size.
- 3. Sensitivity.** It should distinguish ansätze that differ in expressivity, not just in trivial ways.
- 4. Interpretability.** A user should be able to read the number and know what it means qualitatively (“more expressive” or “less expressive”).
- 5. Connection to trainability.** Ideally, the metric correlates with trainability outcomes — knowing the number should help you predict whether the ansatz will train well.

No single metric meets all five. Each one trades off some desiderata for others. The literature has converged on a small handful that are “good enough” for most purposes.

Metric 1: Sim-Johnson-Killoran Expressibility

Already introduced in 8.1 — included here for completeness with the computational details.

Definition. Let $P_{\text{ansatz}}(F)$ be the distribution of fidelities $F = |\langle \psi(\theta_1) | \psi(\theta_2) \rangle|^2$ over uniformly sampled parameter pairs. Let $P_{\text{Haar}}(F) = (d-1)(1-F)^{d-2}$ for $d = 2^n$. Then

$$\text{Expr}_{\text{SJK}} = D_{KL}(P_{\text{ansatz}} \parallel P_{\text{Haar}}).$$

Computational protocol: 1. Sample M parameter pairs $(\theta_1^{(m)}, \theta_2^{(m)})$. 2. For each, compute $F^{(m)} = |\langle \psi(\theta_1^{(m)}) | \psi(\theta_2^{(m)}) \rangle|^2$. 3. Bin the fidelities into B bins to get an empirical distribution $\hat{P}_{\text{ansatz}}$. 4. Compute the KL divergence to the analytic Haar distribution.

Cost: M state preparation pairs, each requiring fidelity computation (which is $O(\text{poly}(n))$ on a simulator and harder on hardware — usually estimated via SWAP test or inversion test).

Interpretation: - $\text{Expr}_{\text{SJK}} \rightarrow 0$: ansatz fidelity distribution matches Haar — fully expressive (Haar-like coverage). - Expr_{SJK} large: ansatz is structured, far from Haar.

Strengths: computable, well-studied, comparable across ansätze with the same n . **Weaknesses:** dominated by global structure; misses fine-grained expressivity differences. The KL is sensitive to bin width and to the number of samples.

Metric 2: Entangling Capability

Also from Sim-Johnson-Killoran (2019). Measures the **average entanglement** of states produced by the ansatz.

Definition. Let $Q(|\psi\rangle)$ be the **Meyer-Wallach Q-measure**:

$$Q(|\psi\rangle) = 2 \left(1 - \frac{1}{n} \sum_{k=1}^n \text{Tr}(\rho_k^2) \right),$$

where ρ_k is the reduced density matrix on qubit k . $Q = 0$ for product states, $Q = 1$ for maximally entangled (in a specific sense). Then

$$\text{Ent}(U) = \mathbb{E}_{\theta}[Q(|\psi(\theta)\rangle)].$$

Computational protocol: sample parameters, prepare states, compute Q for each, average.

Cost: M state preparations and $O(n)$ reduced density matrices each.

Interpretation: - $\text{Ent} \approx 0$: the ansatz produces near-product states. Limited expressivity. - $\text{Ent} \approx 1$: the ansatz produces highly entangled states. - Haar value: roughly $1 - 1/(2^n + 1) \approx 1$ for large n .

Strengths: physically meaningful, easy to compute, complementary to SJK expressibility. **Weaknesses:** entanglement is a *consequence* of expressivity, not the same thing — high entanglement does not imply high reachability.

A good practice is to report **both** SJK expressibility and entangling capability, since they measure different facets.

Metric 3: Dynamical Lie Algebra Dimension

Already sketched in 8.1. A theoretically clean but computationally intensive metric.

Definition. Given the ansatz $U(\theta) = \prod_j e^{-i\theta_j H_j}$, construct the dynamical Lie algebra $\mathfrak{g}$ by closing $\{H_j\}$ under commutators. Then

$$\text{Expr}_{\text{DLA}}(U) = \dim \mathfrak{g}.$$

Computational protocol: symbolic Lie algebra closure. Start with the generators, compute pairwise commutators, add new linearly independent operators, repeat until closure.

Cost: $O(\dim \mathfrak{g}^2)$ commutator computations and linear-independence checks. For an ansatz where $\dim \mathfrak{g}$ is polynomial in n , this is feasible. For deep ansätze where $\dim \mathfrak{g}$ is exponential, intractable.

Interpretation: - Small DLA dimension $\rightarrow$ small reachable manifold $\rightarrow$ low expressivity. - DLA dimension polynomial in $n \rightarrow$ “polynomially expressive” — typical regime for problem-inspired ansätze. - DLA dimension exponential in $n \rightarrow$ “fully expressive” — typical for HEA at sufficient depth.

Strengths: rigorous, ansatz-intrinsic (doesn’t depend on parameter sampling), connects directly to trainability via the Ragone et al. theory. **Weaknesses:** scales poorly with depth; doesn’t capture continuous expressivity differences (it’s a single integer).

Use DLA dimension for theoretical analysis; use SJK expressibility for empirical comparison.

Metric 4: Effective Dimension

Introduced by Abbas et al. (2021) — “The power of quantum neural networks.”

The idea: measure the **information geometric dimension** of the parameter space, weighted by the local geometry of the cost landscape.

Definition. Let $I(\boldsymbol{\theta})$ be the Fisher information matrix at $\boldsymbol{\theta}$. The effective dimension at scale γ is

$$d_{\text{eff}}(\gamma) = \frac{2 \log \left(\frac{1}{V_{\Theta}} \int_{\Theta} \sqrt{\det(I_d + \gamma I(\boldsymbol{\theta}))} d\boldsymbol{\theta} \right)}{\log(\gamma / (2\pi \log n))},$$

where I_d is the identity, V_{Θ} is the parameter-space volume, and γ is a scale parameter. (Don’t be intimidated — this is essentially “average rank of the Fisher information.”)

Interpretation: - d_{eff} close to P (number of parameters): all parameters are independently informative — fully effective ansatz. - $d_{\text{eff}} \ll P$: many parameters are redundant — the effective parameter space is much smaller than P .

Strengths: captures redundancy, connects to generalization theory. **Weaknesses:** computing the Fisher information is expensive; the metric is defined at a scale γ and depends on it.

Effective dimension is the right metric when you suspect your ansatz has many redundant parameters (gauge degrees of freedom, near-degenerate directions). For most problem-inspired ansätze, it’s overkill.

Metric 5: Frame Potential

A more theoretical metric: the **frame potential**.

Definition. For an ansatz $U(\boldsymbol{\theta})$ inducing a distribution on unitaries, the t -th frame potential is

$$\mathcal{F}_t(U) = \mathbb{E}_{\boldsymbol{\theta}_1, \boldsymbol{\theta}_2} [|\text{Tr}(U(\boldsymbol{\theta}_1)U(\boldsymbol{\theta}_2)^\dagger)|^{2t}].$$

For Haar-random unitaries, the frame potential takes a specific value depending on t and $d = 2^n$. An ansatz forms a **t-design** if its frame potential matches the Haar value:

$$\mathcal{F}_t(\text{ansatz}) = \mathcal{F}_t(\text{Haar}).$$

Interpretation: - The ansatz is a 1-design if its first moment matches Haar. - A 2-design if both first and second moments match. - More expressivity = higher t -design.

Strengths: rigorous, connects to the proof of barren plateaus (which assumes 2-designs). **Weaknesses:** needs to be computed at each order t ; only practical for low t .

The frame potential is the *right* metric for connecting to barren plateau theory (Module 4 used it implicitly via “approximate 2-design”). For practical comparison of two ansätze, SJK or entangling capability is easier.

Metric 6: Approximation Error On A Test Suite

A purely operational metric. Define a set of “test states” $\{|\phi_k\rangle\}_{k=1}^K$. For each, compute the minimum error over the ansatz:

$$\epsilon_k = \min_{\boldsymbol{\theta}} \| |\phi_k\rangle - |\psi(\boldsymbol{\theta})\rangle \|^2.$$

Average over the test suite:

$$\text{Expr}_{\text{approx}} = \frac{1}{K} \sum_{k=1}^K \epsilon_k.$$

Interpretation: - Low average error $\rightarrow$ ansatz can approximate diverse target states $\rightarrow$ high expressivity. - High average error $\rightarrow$ ansatz misses many targets $\rightarrow$ low expressivity.

Strengths: directly operational — measures what you actually care about (can the ansatz represent my target?). **Weaknesses:** depends on the test suite. A test suite biased toward your ansatz’s structure will give an unfair advantage.

Use this when you have a specific application in mind and a representative set of target states for that application.

Comparing The Metrics

A summary table:

Metric	Cost	Output range	Captures	Best for
SJK Expressibility	Moderate (samples)	KL divergence (positive)	Distributional coverage	General ansatz comparison
Entangling Capability	Cheap	$[0, 1]$	Entanglement structure	Quick check
DLA Dimension	Expensive (symbolic)	Integer	Reachable manifold dim	Theoretical analysis
Effective Dimension	Expensive	$[0, P]$	Parameter redundancy	Over- parameterized ansätze
Frame Potential	Moderate	Real number	t-design property	Plateau theory
Approximation Error	Application-dep.	Real number	Target representability	Application-driven

For most VQA work, the practical recommendation is: report **SJK expressibility + entangling capability** as the standard pair. Add **DLA dimension** if you can compute it. Add **effective dimension** if your ansatz has visible redundancy.

A Worked Comparison

Consider two 4-qubit ansätze:

Ansatz A: Hardware-Efficient, 2 layers. Single-qubit R_X, R_Y, R_Z on each qubit, then nearest-neighbor CNOTs. Repeat. Total: 24 parameters.

Ansatz B: UCCSD with 4 electrons. Trotterized particle-number-conserving excitations. Total: ~12 parameters.

SJK Expressibility: - A: ≈ 0.05 (close to Haar, highly expressive). - B: ≈ 0.5 (far from Haar — strongly restricted by particle-number conservation).

Entangling Capability: - A: ≈ 0.85 (nearly Haar entanglement). - B: ≈ 0.4 (limited entanglement — restricted to particle-number subspace).

DLA Dimension: - A: full $u(2^4) = 256$ — exponential in n . - B: $\binom{8}{4}^2 = 4900$ but restricted to particle-number sector, so $\binom{8}{4} = 70$.

Effective Dimension at $\gamma = 1$: - A: ≈ 18 (out of 24, some redundancy from gauge freedoms). - B: ≈ 11 (out of 12, near-fully effective).

Approximation Error on the H_2 ground state: - A: ≈ 0.001 (excellent — can represent the ground state). - B: ≈ 0.0001 (even better — UCCSD is *designed* for chemistry).

Conclusion: Ansatz A is more expressive in every general sense, but Ansatz B is *more expressive on the actual problem* because its restricted form happens to align with the chemistry symmetries. Both can solve H_2 ; they just use very different mechanisms.

Caveats On Metric Interpretation

1. Metrics measure structure, not performance. A more expressive ansatz isn't always better. As Module 4 showed, expressivity correlates with trainability obstacles. The right amount is “just enough.”

2. Metric values depend on the parameter range. SJK expressibility, entangling capability, and effective dimension all depend on the parameter sampling distribution. Default is uniform on $[0, 2\pi]^P$, but if you constrain parameters (e.g., $[-\pi/4, \pi/4]$), the metrics change.

3. Sample noise. Empirical metrics (SJK, entangling capability) have sampling noise that scales as $1/\sqrt{M}$. For meaningful comparison, you need M large enough to resolve the difference between ansätze.

4. Comparability across qubit count. SJK expressibility and DLA dimension scale differently with n , so comparing ansätze on different qubit counts requires care. Within a fixed n , comparison is meaningful.

5. No metric captures cost-landscape geometry. None of the standard metrics tells you about local minima, barren plateaus, or convergence rates. They capture *static* expressivity, not *dynamic* trainability.

For trainability-aware ansatz comparison, you need the metrics from node 8.4 in addition.

Structural Role

This node is the **operational toolkit** of Module 8. It enumerates the actual numbers you compute and report when designing or analyzing an ansatz, with practical guidance on which metrics to use for which purpose.

It is the bridge between the abstract definitions of 8.1 and the trainability bounds of 8.4.

Relations to Other Concepts

- **Sim-Johnson-Killoran (2019)** — the original SJK metric paper.
 - **Abbas et al. (2021)** — effective dimension.
 - **Meyer-Wallach Q-measure (2002)** — the entanglement metric.
 - **Frame potentials (Scott 2008)** — t-design theory.
 - **Ragone et al. (2024)** — DLA expressivity.
-

Worked Example — Computing SJK Expressibility For A Toy Ansatz

Consider the 2-qubit ansatz $U(\theta_1, \theta_2) = R_Y(\theta_1) \otimes R_Y(\theta_2)$ — pure product, no entanglement.

Step 1: Sample parameter pairs. Draw $M = 1000$ pairs of (θ_1, θ_2) and (θ'_1, θ'_2) uniformly from $[0, 2\pi]^2$.

Step 2: Compute fidelities. $|\langle \psi(\theta_1, \theta_2) | \psi(\theta'_1, \theta'_2) \rangle|^2 = |\langle 0 | R_Y^\dagger(\theta_1) R_Y(\theta'_1) | 0 \rangle|^2 \cdot |\langle 0 | R_Y^\dagger(\theta_2) R_Y(\theta'_2) | 0 \rangle|^2$.

Each factor is $\cos^2((\theta'_1 - \theta_1)/2)$, uniformly distributed in $[0, 1]$. The product is the product of two independent uniforms — distribution is $-\log(F)$ (concentrated near $F = 0$, heavy near 1).

Step 3: Compare to Haar. For $d = 4$, $P_{\text{Haar}}(F) = 3(1 - F)^2$. The ansatz distribution is dominated by small fidelities (most product states are nearly orthogonal), but the *concentration profile* is very different from Haar.

Step 4: Compute KL. The KL divergence is large (~ 1.5 in natural log units), reflecting that this ansatz is *far* from Haar.

For comparison, a 2-qubit ansatz with one CNOT after the rotations has SJK expressibility ≈ 0.3 — much closer to Haar but still structured.

A deep HEA on 2 qubits has SJK expressibility ≈ 0.01 — essentially Haar.

This worked example shows the kind of empirical comparison that drives ansatz design in practice.

Visualization

Pending. A bar chart of SJK expressibility for 6 standard ansätze (HEA-1, HEA-3, HEA-deep, UCCSD, MPS-2, MPS-8) on 4 qubits, with the Haar value marked as a horizontal line. Caption: “Lower bars are closer to Haar — more expressive.”

Exercises

1. Compute the entangling capability of $U(\theta) = R_Y(\theta) \otimes R_Y(\theta)$ (2 qubits, 1 parameter, no entangling gate). What is it?
 2. The SJK expressibility for an ansatz that produces only the state $|+\rangle^{\otimes n}$ regardless of parameters is what? (Hint: the fidelity distribution is a delta at 1.)
 3. (VQA) For a 6-qubit HEA with 5 layers and a UCCSD ansatz with 6 electrons, predict the relative ordering of (a) SJK expressibility, (b) entangling capability, (c) DLA dimension. State any assumptions.
-

Research Extensions

- **Trainability-aware expressibility metrics** — metrics that incorporate gradient variance, not just static structure.

- **Application-conditioned expressibility** — metrics defined relative to a specific Hamiltonian or target distribution.
 - **Hierarchical metrics** — multi-resolution expressibility (local vs global).
 - **Sample-efficient expressibility estimation** — getting useful numbers from few samples.
-

Cross-links to Other Courses

- **Information Geometry** — Fisher information, effective dimension.
 - **Random Matrix Theory** — Haar measure, frame potentials.
 - **Module 4 (Barren Plateaus)** — the trainability cost of high expressibility.
 - **Module 2 (Ansatz Design)** — the practical context.
-

Variational Quantum Algorithms · Module 8 — Expressivity Vs Trainability · Node 8.2

Concepts: parameterized-circuits, dimensionality, concentration-of-measure

Prerequisites: 8.1

Enables: 8.3, 8.5

hbar.university — own.your.cognition